

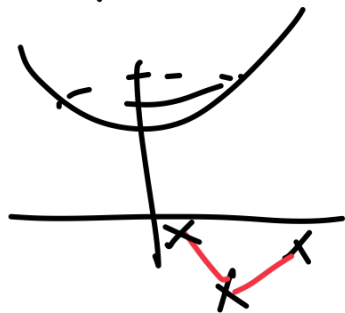
Linear regression $x \in \mathbb{R}^D$

$$h_{\beta}(x) = \beta_0 + \beta_1 x_1 + \dots + \beta_D x_D$$

$$\{x^{(i)}, t^{(i)}\}_{i=1}^N$$

$$\rightarrow l(\beta) = \min_{\beta} \sum_{i=1}^N (t^{(i)} - h_{\beta}(x^{(i)}))^2$$

$\rightarrow$ approach: gradient descent



$$\beta^{(0)} \quad \beta^{(k+1)} \leftarrow \beta^{(k)} - \eta \operatorname{grad}_{\beta} l(\beta)$$

possible extension to datasets not linearly distributed

$$\text{gradient} = \left(\frac{\partial \mathcal{L}}{\partial \beta_0}, \frac{\partial \mathcal{L}}{\partial \beta_1}, \dots, \frac{\partial \mathcal{L}}{\partial \beta_D} \right)$$

$$X = (x_1, x_2, \dots, x_D) \in \mathbb{R}^D$$

$$h_{\beta}(x^{(i)}) = \beta_0 + \beta_1 x_1^{(i)} + \dots + \beta_D x_D^{(i)} \\ = \langle \bar{\beta}, (1, \vec{x}) \rangle$$

$$\tilde{x} = [1, \vec{x}]$$

$$\tilde{x} = [1, x_1, x_2, \dots, x_D]$$

$$h_{\beta}(x^{(i)}) = \langle \bar{\beta}, \tilde{x} \rangle$$

$$= \bar{\beta}^T \tilde{x} = [\beta_0, \beta_1, \dots, \beta_D] \begin{bmatrix} 1 \\ x_1 \\ \vdots \\ x_D \end{bmatrix}$$

$$\min_{\beta} \frac{1}{N} \sum_{i \in \mathcal{D}} (t^{(i)} - h_{\beta}(x^{(i)}))^2$$

(I)

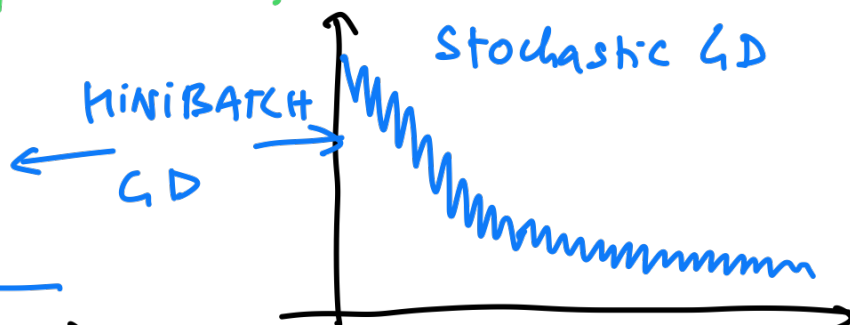
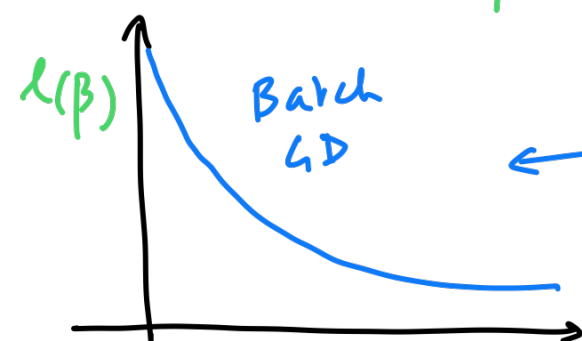
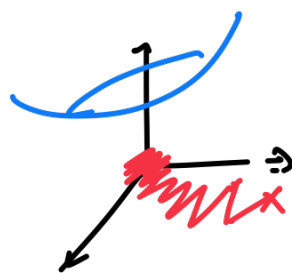
$$\beta^{(k+1)} \leftarrow \beta^{(k)} - \eta \text{grad}_{\beta} \mathcal{L}_{B_{GD}}(\beta)$$

Batch
gradient
descent

(II) for $i \in \mathcal{D}$

$$\min_{\beta} (t^{(i)} - h_{\beta}(x^{(i)}))^2$$

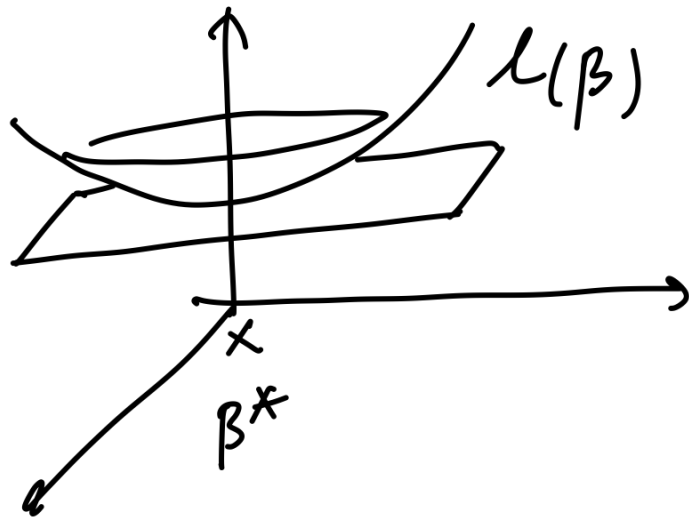
$$\beta^{(k+1)} \leftarrow \beta^{(k)} - \eta \text{grad} \mathcal{L}_{SGD}(\beta)$$



by incorporating polynomial features.

Today: approach #2: Normal equations

statistical intuition



$$\tilde{X} = \begin{bmatrix} \tilde{x}^{(1)} \\ \tilde{x}^{(2)} \\ \vdots \\ \tilde{x}^{(N)} \end{bmatrix}$$
$$\vec{t} = \begin{bmatrix} t^{(1)} \\ \vdots \\ t^{(N)} \end{bmatrix} \quad \vec{\beta} = \begin{bmatrix} \beta_0 \\ \beta_1 \\ \vdots \\ \beta_D \end{bmatrix}$$

$$\vec{e} = \vec{t} - \tilde{X} \vec{\beta} \quad \|\vec{e}\|^2 = e^T e = \langle e, e \rangle$$

$$e = (e^{(1)}, e^{(2)}, \dots, e^{(N)}) = \sum_{i=1}^N (e^{(i)})^2$$

$$e^T e = \sum_{i=1}^N (e^{(i)})^2 = \sum_{i=1}^N (t^{(i)} - (\tilde{x}^{(i)})^T \bar{\beta})^2 = \ell(\beta)$$

$$\ell(\beta) = e^T e = (\bar{t} - \underline{\tilde{X}} \bar{\beta})^T (\bar{t} - \underline{\tilde{X}} \bar{\beta}) \frac{1}{N}$$

$$\ell(\beta) = (\bar{t}^T \bar{t} - \underbrace{\bar{\beta}^T \underline{\tilde{X}}^T \bar{t}}_{-2 \underline{\tilde{X}}^T \bar{t}} - \underbrace{\bar{t}^T \underline{\tilde{X}} \bar{\beta}}_{-2 \underline{\tilde{X}}^T \bar{t}} + \underbrace{\bar{\beta}^T \underline{\tilde{X}}^T \underline{\tilde{X}} \bar{\beta}}_{2 \underline{\tilde{X}}^T \underline{\tilde{X}} \bar{\beta}}) \frac{1}{N}$$

$$\text{grad}_{\beta} \beta^T A \beta = \text{grad}_{\beta} [\beta_0, \beta_1] \begin{bmatrix} a_{11} & a_{12} \\ a_{12} & a_{21} \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \end{bmatrix} = 2A \beta$$

$$\text{grad}_{\beta} \quad \beta_0^2 a_{11} + \beta_1^2 a_{22} + \beta_0 \beta_1 a_{12} \\ + \beta_1 \beta_0 a_{12}$$

$$= \beta_0^2 a_{11} + \beta_1^2 a_{22} + 2\beta_0 \beta_1 a_{12}$$

$$\text{grad}_{\beta} = \left[\frac{\partial \mathcal{L}}{\partial \beta_0}, \frac{\partial \mathcal{L}}{\partial \beta_1} \right] = \begin{bmatrix} 2\beta_0 a_{11} + 2\beta_1 a_{12}, \\ 2\beta_1 a_{22} + 2\beta_0 a_{12} \end{bmatrix}$$

$$2 \begin{bmatrix} a_{11} & a_{12} \\ a_{12} & a_{22} \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1 \end{bmatrix}$$

$$\text{grad}_{\beta} \mathcal{L}(\beta) = -2 \underline{\tilde{X}}^T \underline{t} + 2 \underline{\tilde{X}}^T \underline{\tilde{X}} \beta = 0$$

$$\underline{\tilde{X}}^T \underline{\tilde{X}} \beta = \underline{\tilde{X}}^T \underline{t} \rightarrow \text{Normal equations}$$

$$\begin{bmatrix} 1 & & & 0 \\ & \ddots & & \\ 0 & & & 1 \end{bmatrix}$$

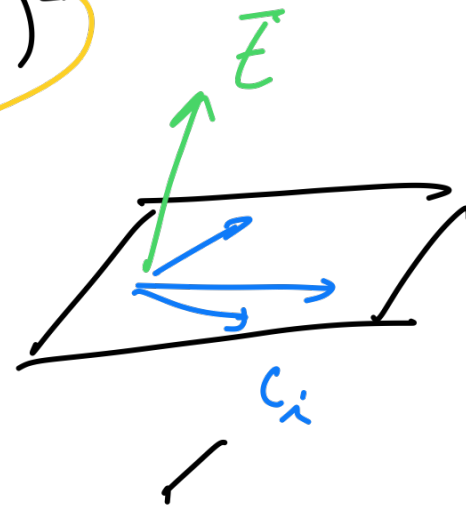
$$\underbrace{(\tilde{\underline{X}}^T \tilde{\underline{X}})^{-1} (\tilde{\underline{X}}^T \tilde{\underline{X}})}_{\mathbf{I}} \beta = (\tilde{\underline{X}}^T \tilde{\underline{X}})^{-1} \tilde{\underline{X}}^T \underline{\underline{e}}$$

$$\underline{\underline{\beta}} = (\tilde{\underline{X}}^T \tilde{\underline{X}})^{-1} \tilde{\underline{X}}^T \underline{\underline{e}}$$

$$\min_{\beta} \| \underline{t} - \tilde{\underline{X}} \beta \|_2^2 = \sum_{i=1}^N (t^{(i)} - \tilde{\underline{x}}^{(i)T} \beta)^2$$

$$\tilde{\underline{X}} = \begin{bmatrix} \text{---} x^{(1)} \text{---} \\ \text{---} x^{(2)} \text{---} \\ \vdots \\ \text{---} x^{(N)} \text{---} \end{bmatrix}$$

$$\beta = \underline{t}$$



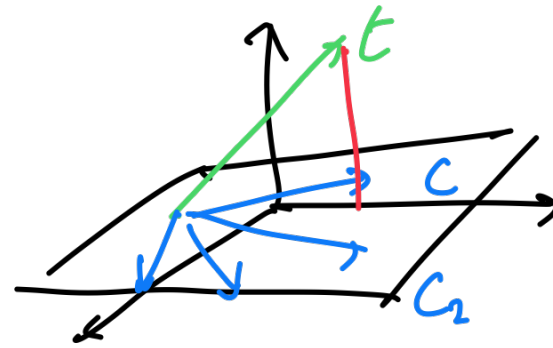
columns of $\tilde{\underline{X}}$

$$\| \underline{e} \|_2^2 =$$

$$\sum_{i=1}^N (e^{(i)})^2$$

Condition to use the normal equations $\tilde{X}^T \tilde{X}$ has to be invertible

$$\tilde{X} = \begin{bmatrix} | & | & | & \dots & | \\ c_0 & c_1 & c_2 & \dots & c_p \\ | & | & | & \dots & | \end{bmatrix}$$



What happens if c_k can be (almost) derived from the other columns c_i

When solving $\min_{\beta} \sum_{i=1}^N (t^{(i)} - (\tilde{X}^{(i)})^T \beta)^2 = \|\underbrace{\vec{t} - \tilde{X} \beta}_{\text{residual}}\|_2^2$

there still exists an orthogonal projection but this projection is not uniquely defined from the c_i 's

Assume columns C_i 's of $\tilde{X}$ are almost linearly dependent

Assume we want to write $\vec{e}$ from the first p columns:

Start with $z_0 = C_0$

$$\langle z_e, z_e \rangle = z_e^T z_e$$

for $j = 1, \dots, p$

$$\hat{\gamma}_{lj} = \langle C_j, z_l \rangle / \langle z_l, z_l \rangle \quad \text{for } l = 0, \dots, j-1$$

$$z_j = \underbrace{C_j} - \sum_{k=0}^{j-1} \hat{\gamma}_{kj} \underbrace{z_k}$$

z_p is the only column that contains column C_p and with a coefficient 1

$\Rightarrow$ as a result $\hat{\beta}_p$ can be obtained by projecting $\vec{t}$ onto z_p

$$\hat{\beta}_p = \frac{\langle z_p, \vec{t} \rangle}{\langle z_p, z_p \rangle} \rightarrow \ll \epsilon$$

$$\hat{\beta}_p \gg \epsilon$$

For large regression coefficients

The model $h_{\beta}(x) = \beta_0 + \beta_1 x_1 + \dots + \hat{\beta}_p x_p + \dots$

applied to a new x will result in

inaccuracies (small changes in x_p will lead

to very different predictions)

→ Next Wednesday

→ alternative approaches

→ Best subset selection

→ Ridge Regression

→ LASSO